



Generic Comparisons

Citation

Nickel, Bernhard. 2010. Generic comparisons. *Journal of Semantics* 27(2): 207-242.

Published Version

doi:10.1093/jos/ffq004

Permanent link

<http://nrs.harvard.edu/urn-3:HUL.InstRepos:4692277>

Terms of Use

This article was downloaded from Harvard University's DASH repository, and is made available under the terms and conditions applicable to Open Access Policy Articles, as set forth at <http://nrs.harvard.edu/urn-3:HUL.InstRepos:dash.current.terms-of-use#OAP>

Share Your Story

The Harvard community has made this article openly available.
Please share how this access benefits you. [Submit a story](#).

[Accessibility](#)

Generic Comparisons*

Bernhard Nickel · Harvard University

January 17, 2010

Abstract

The paper discusses comparative generic sentences *As are F-er than Bs*—*girls do better than boys in grade school*, for example—which pose severe problems for extant accounts. In their stead, the paper proposes reconceiving the LF of generic sentences as more closely akin to that of sentences containing non-generic plurals, paradigmatically plural definite descriptions. Given this one crucial change, several otherwise puzzling features of comparative generics are immediately explicable, including their relatively weak truth conditions and some of the logical relations they enter into.

Keywords SEMANTICS · NATURAL LANGUAGE · GENERICS · COMPARATIVES · PLURALS · DISTRIBUTIVITY

1 Introduction

Many paradigmatic generic sentences seem to express quite strong generalizations, ones that might not hold universally, but whose exceptions are plausibly considered deviations from the norm, such as (1) and (2).

(1) Ravens are black.

(2) Tigers have stripes.

*The initial impetus for writing this paper came from a challenge by Carrie Jenkins to account for generic comparisons within the framework of quantificational theories of generics. An ancestor of this paper was given at the conference ‘Generics: Interpretation and Use’ in Paris in 2009. I want to thank the participants there for many helpful suggestions. I also want to thank three anonymous reviewers for this journal.

A semantic theory for generics starting out from these examples may well analyze them as true just in case most or all of the normal members of the kind satisfy the property predicated. Let's call such a theory a strong quantificational theory—I'll give a more precise statement below.

There are, however, many examples that cast doubt on the basic workability of this approach. These are all true generics which an analysis in terms of *all normal* or *most* seems to predict to be false. This paper discusses a subclass of these problem cases, illustrated by (3) and (4).

(3) Girls do better than boys in grade school.

(4) Horses are taller than cows.

(3) is not appropriately paraphrased as saying that all normal girls do better in grade school than all normal boys nor as saying that most girls do better than most boys. The former is clearly too strong, the latter can be true even when (3) is false. The same holds, *mutatis mutandis*, of (4). Sentences like (3) and (4) are what I'll call *generic comparisons*: sentences of the form *As are F-er than Bs*.

Generic comparisons are particularly interesting for at least two reasons. First, they resist a treatment that is plausible for many other potential counterexamples to strong quantificational theories. *Lions have manes* or *chickens lay eggs* are true, even though neither all normal nor most lions have manes, and even though neither all normal nor most chickens lay eggs. These cases can nonetheless be analyzed as involving a suitable restriction of the domain of quantification so that, in the restricted domain, the paraphrase in terms of *most* or *all normal* is appropriate. Proponents of this strategy have to answer many questions, such as what exactly the relevant restriction is and what factors determine it. But at least it does not seem completely *ad hoc* that such a restriction is available. Unfortunately, appeals to such a restriction won't help with generic comparisons. Very basic features of strong quantificational theories conspire to predict truth conditions for generic comparatives that are far too strong.

Second, these examples are completely systematic. Other sentences that are often cited as problematic for strong quantificational theories, such as *mosquitos carry plasmodia* or *sharks attack bathers* are isolated, and the intuitions of well-informed native speakers vary on their acceptability. By con-

trast, generic comparisons form a systematic class of troublesome cases, and intuitions concerning their truth-values are very firm.

In particular, I want to highlight two relevant cases that will guide much of the discussion to follow. First, here is a case that intuitively verifies (3), described in terms of a distribution of scholastic achievement (1)—I’ll call it the **SHIFT** scenario.

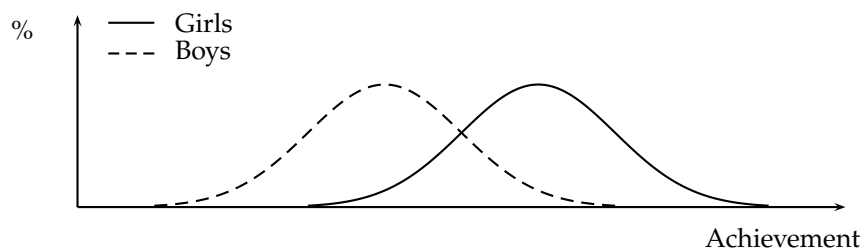


Figure 1: **SHIFT**

Though some boys outperform many of the girls, the distribution of scholastic achievement among girls is systematically shifted to the right of that among boys. By contrast, here is a situation that intuitively falsifies (3). I’ll call it the **SANDWICH** scenario.

[**SANDWICH**] Scholastic achievement among boys is quite heavily polarized: a third of all boys are excellent, so good in fact that the boys in that third are all better than any girl. The other two thirds of boys is terrible, so bad in fact that the boys in these two thirds are all worse than even the weakest girl.

The challenge for semantic theories of generics is to predict that (3) is true in the **SHIFT** but false in the **SANDWICH** scenario.¹ One of my major claims in this paper is that no extant theory can meet it.

Several options are available to attempt a solution to the problem, drawing on one or another proposal offered to deal with other problem cases. Below, I’ll consider the prospects of analyzing generic comparisons as kind predications, following some remarks of Krifka et al. (1995, 83) who recommend this strategy for many otherwise troublesome examples. I’ll also consider the

¹Note, for example, that a paraphrase of (3) in terms of *most* is true in the **SANDWICH** scenario.

possibility of using the alternative generic quantifier that Cohen (1999b, 54ff) introduces to deal with the (in)famous (5).

- (5) Dutchmen are good sailors.

I'll argue that both of these approaches are empirically inadequate, as are some natural extensions of their motivating ideas not discussed in the literature. In their stead, I'll propose a compositional semantics for generic comparisons that derives their relatively weak truth conditions as an interaction effect between a strong generic quantificational element and independently motivated aspects of the interpretation of comparatives in plural constructions. Doing so forces us to reconceive the LF of even simple generic sentences like (1) and (2).

After some preliminaries in §2, I formulate the problem the interpretation of generic comparisons pose more precisely in §3. There, I also explain why I reject alternative approaches. §4 contains the positive proposal, beginning with a discussion of comparatives in non-generic plurals and extending the treatment to generic comparisons.

2 Preliminaries

I have so far spoken of generics *simpliciter*, but I need to narrow my focus since generics form a semantically and syntactically heterogeneous class. Some seem to explicitly predicate a property of a kind, such as *ravens are widespread* or *dodos are extinct*. I want to set these aside.² I also want to set aside merely existential sentences containing bare plurals, such as *ravens are on my lawn*. Finally, I'll only discuss sentences with bare plural subjects, not ones with singular definite or indefinite descriptions. I'll call the sentences under investigation *characterizing sentences*.³

I impose these restrictions because the phenomenon of interest seems to be confined to bare plural generics, as the examples in (6) illustrate.

- (6) a. A girl does better than a boy in grade school.

²That is not to say that a semantics for generics can simply treat the fact that bare plurals appear in both kind-predications and in generalizations like (3) and (4) as an orthographic accident. But we face enough problems without trying to solve that one, too.

³Save for my restriction to bare plurals, my terminology coincides with that of Krifka et al. (1995).

- b. The girl does better than the boy in grade school.

While (3) can be true in a situation in which there is significant overlap in performance, such as the *SHIFT*, (6a) and (6b) cannot—if they are even interpretable as generics.

A *quantificational* account of characterizing sentences adopts two hypotheses, one about the proposition these sentences express, the other about their LF and compositional semantics. The first holds that there is a systematic, albeit complex, relationship between facts about the distribution of properties among individuals (perhaps across possible worlds and times) and the truth of a distinctively generic proposition conveyed by a characterizing sentence. It holds, for example, that there are certain distributions of blackness among ravens (across space, time, and possible worlds) such that, if and only if any one of them obtains, the generic proposition that ravens are black is true. Likewise, there are certain distributions of stripiness among tigers (across space, time, and possible worlds) such that, if and only if any one of them obtains, the generic proposition that tigers have stripes is true. We can mark this correlation between facts about individuals and the truth of generic propositions by introducing a generic quantifier *GEN* into our metalanguage and saying that the generic proposition expressed by, e.g., *ravens are black* is a quantified proposition, the proposition [*GEN* x : *Raven*(x)](*Black*(x)).

At this point, we face the theoretically important question about how to further describe this proposition, which is to say, what the systematic relationship between facts about individuals and the truth of generic propositions is. Much of the debate concerning the semantics of characterizing sentences is about just this question, although strictly speaking, this is not directly a semantic issue.⁴ It is not a semantic issue because so long as we agree that generics express propositions, all of the prevailing semantic approaches are compatible with the view that generic facts that make generic propositions true are not metaphysically basic, including kind-predicating views such as Carlson (1977).

The second hypothesis characteristic of quantificational accounts specifically concerns semantics. Not only is the proposition expressed by a characterizing sentence quantificational, its logical form is, too. A quantificational

⁴See Krifka et al. (1995, §1.2.6.) for an overview of some of the options, as well as Pelletier and Asher (1997).

element in their LF, call it *gen*, encodes the relationship between facts about individuals and the truth of the generic proposition expressed by the sentence.⁵ Saying only this much about quantificational theories leaves open various semantic and syntactic options. They leave open where in the LF the generic quantifier originates, whether it is a determiner of the bare plural, an adverb of quantification, or something else. They also leave open the meaning of *gen*.

A *strong* quantificational account takes a stand on that meaning. Since generics tolerate exceptions, *gen* can obviously not be analyzed as a universal quantifier. A strong quantificational account analyzes it as something that falls just short, such as quantification over most of the members of the kind or all of the normal ones. I already mentioned the intuitive sense that the initial examples (1) and (2) are appropriately paraphrased as strong generalizations. Perhaps more important is the more directly empirical observation that some generics are false even though most members of the kind at issue conform to the generalization, as in (7).

- (7) a. Books are paperbacks.
- b. Prime numbers are odd.

Since (7a) and (7b) are both false, even though the majority of books are paperbacks, and the vast majority of prime numbers are odd, (7a) and (7b) must express very strong generalizations.⁶

3 Inappropriate Truth Conditions

Consider again (3), repeated here.

- (3) Girls do better than boys in grade school.

⁵Thus, *gen* is the element in the object language that appears in the LF of characterizing sentences, GEN the quantifier in the metalanguage used to interpret the object language *gen*.

⁶Quantifiers in the generalized quantifier description are usually identified as second-order properties that satisfy permutation invariance, conservativity, and extension. However, for the purposes of this discussion, I won't assume that the generic operator GEN satisfies these constraints—this practice seems to be in line with how quantificational accounts of generics are usually discussed. The reason to call the accounts I discuss in the main text *quantificational* is that there is an element that encodes the relationship between the instantiation of a property among the individuals making up a kind and the corresponding generic, a relationship that is broadly speaking about *how many*.

We need make only very few and weak assumptions in order to predict inappropriate truth conditions for (3). Being explicit about them allows us to categorize responses to the problem posed.

First, we need to assume that (3) has a quantificational LF. Second, both bare plurals are interpreted generically, so that both are associated with a generic quantifier. Third, the generic quantifier that appears in the LF of (3) is the same strong quantifier as appears in the LF of the initial examples (1) and (2). Fourth, and finally, the generic quantifier occupies one of two standard positions for quantifiers: a nominal determiner or an unselective binder modeled on Lewis (1973).⁷

At this point, it will be useful to say something about how I will assume we interpret comparatives, and especially how quantifiers are interpreted when they occur in the scope of comparatives. Since I want to argue that we cannot get a proper treatment of generic comparisons given the assumptions I've just set out, I want to assume the most flexible theory of comparatives and their interactions with quantification. I will essentially employ the theory offered by Schwarzschild (2008), what he calls the *A-not-A* approach. On this view, a comparative *a is F-er than b* is interpreted in two steps. First, *a* and *b* are mapped to a degree on a scale associated with *F*-ness. The comparative then says that *a* meets or exceeds some threshold that *b* fails to meet or exceed (i.e., one that *b* falls below).⁸ Consider example (8).

(8) John is taller than Bill.

In order to evaluate this sentence, John and Bill are both mapped to a degree of height (the scale associated with tallness), and (8) is true iff there is some threshold that John's degree of height meets or exceeds and that Bill's height neither meets nor exceeds. In formal terms:

(9) $\exists \theta (\text{height}(j, \theta) \wedge \neg \text{height}(b, \theta))$

Schwarzschild's proposal is especially well-suited to my purposes since it posits that the predicate appears twice in the semantic interpretation of the comparative, and the occurrence that corresponds to the object in the initial sentence is within the scope of negation. That means that once we introduce

⁷For the determiner option, see, e.g., Asher and Morreau (1995); for the adverb of quantification option, see Cohen (1999a,b); Schubert and Pelletier (1989); Wilkinson (1991).

⁸See, also Kennedy (1999, 2007).

other scope-taking operators, such as quantifiers, we have further options to consider. But if extant treatments of generics still cannot give a proper interpretation of (3), that gives us good reason to look elsewhere.⁹

With this in mind, let me turn to the interpretation of (3). By the first assumption, (3) has quantificational truth conditions. By the second, both bare plurals introduce variables bound by a generic quantifier. Suppose that this generic quantifier is a nominal determiner (the first option for the fourth assumption). In that case, (3) can be paraphrased as saying that GEN-many girls do better in grade school than GEN-many boys. Put in formal terms, we have two possible interpretations here, depending on how the scope of generic quantifier that binds the variable associated with *boys* is related to the negation in the comparative construction. The possibilities are given in (10a) and (10b).

- (10) a. $[\text{GEN } x: \text{Girl}(x)] [\exists \theta] (\text{Does.Well}(x, \theta) \wedge$
 $[\text{GEN } y: \text{Boy}(y)] (\neg \text{Does.Well}(y, \theta)))$
 b. $[\text{GEN } x: \text{Girl}(x)] [\exists \theta] (\text{Does.Well}(x, \theta) \wedge$
 $\neg [\text{GEN } y: \text{Boy}(y)] (\text{Does.Well}(y, \theta)))$

By the third assumption, these formulae have roughly the following truth conditions. We interpret (10a) as saying that all normal or most girls meet or exceed a threshold that all normal or most boys fail to meet. These are the truth conditions most naturally understood for (11a) or (11b).

- (11) a. Every normal girl does better in grade school than every normal boy.
 b. Most girls do better in grade school than most boys.

(11a) is far too strong, since it is false in the SHIFT scenario. (11b) is too weak, since it predicts that (3) is true in the SANDWICH scenario, since all of the girls do better than two thirds of the boys.

Turning now to (10b), we find excessively weak truth conditions. In line with the third assumption, it is interpreted roughly as saying that all normal

⁹This proposal is useful for another reason. It essentially assumes that all comparatives are “clausal,” i.e., that the LF of a sentence like (8), *John is taller than Bill* is *John is taller than Bill is ~~tall~~*. This is a very plausible hypothesis insofar as bare plurals in object position are not usually interpreted generically—see Diesing (1992). That suggests in turn that when bare plurals are interpreted generically, they’re not really in object-position, but are instead in a clause.

or most girls meet or exceed a threshold that not all normal or most boys meet or exceed, i.e., that all normal or most girls exceed a threshold that at least one normal boy fails to meet. In other words, it says that all normal girls do better than the weakest normal boy (and *mutatis mutandis* for a *most*-interpretation). That is far too weak.

The situation is no different if *gen* is an unselective quantifier that occupies the position of an adverb of quantification. In that case, we would paraphrase (3) as saying that generically (generally, typically), a girl does better in grade school than a boy, i.e., as saying that GEN-many pairs $\langle x, y \rangle$ satisfy the condition. Given, again, that we're interpreting the generic quantifier as a relatively strong quantifier, we predict that (3) has roughly the truth conditions of the examples in (12).

- (12) a. In all normal cases, a girl does better than a boy in grade school.
- b. Mostly, a girl does better than a boy in grade school.

As before, these truth conditions are far too strong. Regardless of whether we pursue the nominal determiner or the adverb of quantification options, the asserted relation—doing better in grade school—has to hold between any pair we can form by taking one of GEN-many girls and one of GEN-many boys, and that in turn entails that the weakest of the GEN-many girls still does better than the strongest of the GEN-many boys.

One of the four assumptions has to go. My own account will reject the fourth, that the generic quantifier occupies one of the standard positions. By way of motivation, I'll consider rejecting one of the other three. I'll begin with the first, that generic comparisons like (3) should be analyzed quantificationally (§3.1). I'll then consider a proposal on which we reject the third, that the generic quantifier is analyzed as a strong quantifier (§3.2). Finally, I'll consider rejecting the second assumption, that both of the NPs are interpreted generically. The core idea is that the object NP—*boys* in *girls do better than boys in grade school*—is a dependent plural (§3.3).

3.1 Kind-Reference and Other Aggregative Proposals

Generic comparisons are at least superficially similar to paradigmatic characterizing sentences. Nonetheless, one could try to assimilate them to such

kind-predicating sentences as *ravens are widespread*, which predicate a property of a kind directly. This is the proposal of Krifka et al. for many sentences that would otherwise spell trouble for a strong quantificational theory. In essence, the strategy is to transfer the semantics for comparatives involving individuals, such as (8), *John is taller than Bill*, without the detour through quantification. We simply move from comparing particular objects to comparing kinds.

In this case, each of the kinds mentioned is mapped to a degree on an appropriate scale, and the generic comparison *As are F-er than Bs* is true just in case there is a threshold on the scale of *F*-ness that the *As*—considered as a kind—meet or exceed and which the *Bs* fail to meet or exceed. Since it is surely a context-sensitive matter how the relevant degrees are determined, the account makes no general prediction about the relationship between the degrees to which individual *As* and *Bs* are *F* and the degrees the corresponding kinds are assigned. The account therefore isn't saddled with predicting inappropriate truth conditions, largely because it makes no general predictions on this point, at all. But that need not be *ad hoc*. The kind-referring strategy should liken the determination of the relevant mapping to the determination of the reference of demonstratives: saying how either determination is made is no part of the semantics of these expressions, and in neither case should we think that this division of labor is at all problematic.

Prima facie, this proposal is theoretically unattractive because it denies the apparent similarities between generic comparisons and paradigmatic characterizing sentences, since the latter are given a broadly quantificational treatment while the former are not. However, I here only want to focus on a more straightforwardly empirical problem, to wit, that this proposal draws the analogy between comparisons of individual objects and generic comparisons too closely. The two kinds of comparisons exhibit different patterns of logical compatibility and entailment, and the kind-predicating strategy cannot account for this difference.

When comparing two individual objects where both can be associated with degrees on the relevant scale, it will always be true that one stands in the comparative relation to the other, the other to the one, or that they are equal. That is to say, one of the schemas in (13) must be true.¹⁰

¹⁰In example (13c), and all of the corresponding examples below, I insert the parenthetical

- (13) a. *a* is *F*-er than *b*.
 b. *b* is *F*-er than *a*.
 c. *a* is (exactly) as *F* as *b*.

In the case of John, Bill, and comparisons of tallness, if there is a degree to which John is tall, and there is a degree to which Bill is tall, then one of the following three sentences must be true.

- (14) a. John is taller than Bill.
 b. Bill is taller than John.
 c. John is (exactly) as tall as Bill.

However, the corresponding pattern does not hold for generic comparisons, as the examples in (15) show—assume as before that we measure the height of a quadruped at the shoulder.

- (15) a. Cows are taller than horses.
 b. Horses are taller than cows.
 c. Cows are (exactly) as tall as horses.

All of these are false in the actual world—the tallest horses are taller than the tallest cows, and the shortest horses are shorter than the shortest cows. Moreover, the problem is not that cows and horses somehow resist comparison with respect to height, as the truth of the examples in (16) shows.

- (16) a. Cows are taller than cats.
 b. Elephants are taller than horses.

exactly to ensure that we do not read it as *a* is at least as tall as *b*. I do not mean to imply that anything turns on a particular standard of precision. I am also taking for granted that the interpretation of the predicate *F* is the same in all of the examples in (13). Some predicates can be associated with different scales in different contexts, such as *big* or *good*. The reason to require that these predicates are interpreted uniformly can be brought out with an example. Suppose that John is taller and lighter than Bill. In that case, all of the following are false if we consider the dimension of comparison indicated in parentheses.

- (i) a. John is bigger (by weight) than Bill.
 b. Bill is bigger (by height) than John.
 c. John is (exactly) as big (by weight) as Bill.

But so long as we hold the dimension fixed, the schemas in (13) jointly exhaust logical space and one of them must be true.

More generally, then, in the case of generic comparisons, the schemas in (17), corresponding to (13), do not jointly exhaust logical space.

- (17) a. *As* are *F*-er than *Bs*.
 b. *Bs* are *F*-er than *As*.
 c. *As* are (exactly) as *F* as *Bs*.

A proponent of the kind-predicating strategy thus needs to identify some difference between comparisons of individuals and generic comparisons. One initially plausible attempt points to vagueness: kinds cannot be assigned to precise degrees on an associated scale, but only a vague range on such a scale. And we know that, when vagueness is involved, sets of sentences that one might think jointly cover all of logical space (so that at least one of them must be true) can nonetheless all be false. If John is a borderline case of baldness, then it might be false both that he is bald and that he is not.¹¹

But the approach I'm considering on behalf of the kind-predicating strategy is unpromising if we take the appeal to vagueness seriously. It is a commonplace that predicates that exhibit vagueness in their positive form do not exhibit vagueness once they are in an explicitly comparative construction. Two paint chips might both be on the borderline between orange and red, so that neither (18a) nor (18b) are definitely true or definitely false.

- (18) a. Chip₁ is red.
 b. Chip₂ is red.

Even in this situation, the relevant instances of (13), given here in (19) exhaust logical space and one of them is definitely true.

- (19) a. Chip₁ is redder than chip₂.
 b. Chip₂ is redder than chip₁.
 c. Chip₁ is (exactly) as red as chip₂.

The proponent of the kind-predicating strategy should therefore not appeal to vagueness as the model by which to explain the possibility that all of the examples in (17) can be false at the same time.

¹¹Giving up on classical logic in this way is one way of dealing with the initial phenomenon, which is just that we find neither *John is bald* nor *John is not bald* to be unproblematically acceptable. For an overview of this and other options, see Williamson (1994).

Instead, she might point to imprecision as a feature of a situation that *is* preserved even once we move to explicit comparisons and that might account for the observation about (17). However, imprecision generally does not interfere with the usual logical relations. Suppose it is hard to tell exactly how tall John and Bill are. We might model this by mapping them not to a precise degree of height, but to a range of the scale. Even in that case, we are usually happy to accept that one of the sentences in (14)—the height comparisons between John and Bill—is true, even if we don’t know which it is because our measuring situation is unfavorable. Which sentence is true may well also depend on subtle features of the context, but given standards of precision in a context, one of them will be true.

The kind-predicating strategy is therefore unpromising. An appeal to vagueness is untenable because vagueness disappears in comparatives. An appeal to imprecision may be theoretically acceptable, but it does not allow us to predict that all of the sentences are false. Thus, the logical relations that generic comparisons enter into differ from those that comparisons involving individual objects enter into, and the kind-predicating strategy cannot account for this fact.

The problem I have raised for the kind-predicating strategy is a problem for *any* strategy that seeks to aggregate all of the members of the kind that the generic is about, assign that aggregate a single degree, and then compare the degree assigned to one aggregate with that assigned to another aggregate. For example, one might interpret a generic comparison by paraphrasing it in terms of averages. On this approach, (3) is paraphrased as *the average girl does better than the average boy in grade school*.¹² This approach makes precisely the same predictions as the kind-predicating strategy in that comparisons involving averages pattern with comparisons between individual objects, not generic comparisons. In general, one of the schemas in (20) must be true if the average *A* and the average *B* can be assigned any degree of *F*-ness at all.

- (20) a. The average *A* is *F*-er than the average *B*.
 b. The average *B* is *F*-er than the average *A*.
 c. The average *A* is (exactly) as *F* as the average *B*.

¹²Moreover, there are already detailed semantic proposals for *average*, so that a proponent of the aggregative approach can simply help herself to them. See, e.g., Carlson and Pelletier (2002); Kennedy and Stanley (forthcoming).

Hence, any aggregative proposal that seeks to assign a single degree to all of the members of the kind relevant to determining the truth of generic comparisons is inadequate.

3.2 An Alternative Generic Quantifier

I now turn to a way of rejecting the assumption that the quantifier that appears in generic comparisons is the usual strong one. This section focuses on Cohen's introduction of a so-called *relative* generic quantifier. The discussion is largely exploratory, since he is primarily concerned with sentences like (5), *Dutchmen are good sailors*, and does not discuss generic comparisons.¹³ Nonetheless, I want to investigate whether we can extend his treatment to comparative constructions because of the close semantic connection between gradable predicates in their positive form, such as *are good sailors*, and comparatives.

To introduce Cohen's system, let me begin with the core examples (1) and (2), which he calls *absolute* generics. He interprets the generic quantifier in such a way that its restrictor is determined, at least in part, by the predicate via its association with a set of alternatives. To interpret *As are F*, we have to compute the set of alternatives $ALT(F)$. In most cases, F is included in $ALT(F)$, and in most cases, the alternatives are mutually exclusive. For example, to interpret (1), *ravens are black*, we associate the property of being black with alternative colors. With that set in hand, Cohen gives the following truth conditions.¹⁴

- (21) *As are F* is true iff the probability that a randomly chosen A that also satisfies at least one of the properties in $ALT(F)$ is F is greater than .5.

The reason to introduce alternatives is to solve the problem posed by sentences such as *lions have manes*. If we interpreted this simply as saying that the probability that randomly chosen lion has a mane is greater than .5, we'd be committed to more than half of all lions having manes, which would mean that we predict the sentence to be false. The set of alternatives $ALT(F)$ effectively restricts the domain to lions with some form of ornamentation. However, Cohen accepts that his theory as applied (5) yields the truth conditions that the

¹³Except for a passing remark at Cohen (2004, 549).

¹⁴See Cohen (1999b, p. 37).

probability that a randomly picked Dutch sailor is a good one is greater than .5, and that is too strong.

In response, Cohen introduces *relative generics*, generic sentences that are analyzed in terms of an alternative generic operator. Characterizing sentences are therefore systematically ambiguous, depending on whether they are analyzed as absolute or relative. When *As are F* is analyzed as a relative generic, we do not just consider the alternatives to *F*, but also the alternatives to *A*, $\text{ALT}(A)$. In the case of (5), $\text{ALT}(A)$ might include other nationalities. Relative generics have the truth conditions in (22).¹⁵

- (22) *As are F* is true iff the probability that a randomly chosen *A* that satisfies one of the alternatives in $\text{ALT}(A)$ is *F* is higher than the probability that an arbitrarily chosen object that satisfies one of the members of $\text{ALT}(A)$ and one of the members of $\text{ALT}(F)$ is *F*.

As applied to (5), this theory predicts the truth conditions that an arbitrarily chosen Dutch sailor is more likely to be a good sailor than an arbitrarily chosen sailor from one of the alternative nations. We can see why this interpretation is aptly called *relative*. A relative generic requires for its truth that the relevant members of the kind be more likely to satisfy the predicate than members of some other kind(s), i.e., we are relating different kinds, such as Dutchmen and Swiss.

So that we may evaluate this proposal, let me say how claims about probabilities are related to facts “closer to the ground.” For our purposes, we can simply translate talk of probabilities into talk of ratios. To say that the probability that a randomly picked Dutch sailor is good is higher than the probability that a randomly picked sailor from some other country is amounts to the claim that the ratio of good Dutch sailors to Dutch sailors of any skill is higher than the ratio of good sailors from other countries to sailors from these countries of any skill.¹⁶

¹⁵See Cohen (1999b, 55f). One benefit of Cohen’s strategy is that it makes the difference between relative and absolute generics not completely *ad hoc*. As he argues in Cohen (2001), the difference between the two readings can be reduced to a difference in the setting of one parameter, one we also see in some non-generic cases involving *many* and *often*.

¹⁶Within the context of Cohen’s theory, the initial interpretation of generics in terms of probabilities plays other roles than simply introducing ratios. It also allows him to motivate various constraints on the classes within which the relevant ratios are assessed, what he calls homogeneity constraints. For more discussion, see Cohen (1999a, 2004).

Let me now consider whether we can make use of a relative generic quantifier to yield adequate semantics for generic comparisons. The most direct application of the theory analyzes (3) as saying that a randomly picked girl is more likely than a randomly picked member of $\text{ALT}(\llbracket \text{girls} \rrbracket)$ to satisfy the predicate *does well in grade school*. Given the context, especially the linguistic context, it's clear that $\text{ALT}(\llbracket \text{girls} \rrbracket)$ just consists of the set of boys. (3) is thus predicted to have the truth conditions (23).

- (23) A randomly picked girl is more likely to do well in grade school than a randomly picked boy.

As a sidenote, it is somewhat opaque how we could arrive at these truth conditions compositionally. The question really turns on the interpretation of the predicate *does better than boys in grade school*, and specifically *boys*. I've already suggested that we cannot interpret it as an ordinary generic bare plural, which in Cohen's system is interpreted as an absolute generic. That would make the reading too strong. And it is not immediately obvious what a relative reading of that NP would amount to.

The best strategy I can imagine interprets the comparative morphology and the comparative phrase as arguments of the relative generic operator directly. On this approach, the explicitly comparative generic (3) just supplies the arguments explicitly that the positive (relative) generic has to take implicitly. For the particular case of (3), *-er than boys* simply specifies that $\text{ALT}(\llbracket \text{girls} \rrbracket)$ contains the property of being a boy.

But I don't want to focus on the compositional semantics. Even taking for granted that we can predict (23) as stating the truth conditions of (3), there are problems with the proposal. It predicts that a comparative generic has the same truth conditions as a positive relative generic, so long as the same set of alternatives are salient in the context. This prediction fails, as the examples in (24) illustrate.

- (24) a. Feathers are heavier than appleseeds, though of course no feathers nor any appleseeds are heavy.
 b. Molecules are larger than atoms, though of course no molecules nor any atoms are large.
 c. Dwarves are taller than hobbits, though of course no dwarves nor any hobbits are tall.

Each of these examples is predicted to be contradictory on the relative generic strategy. Consider (24a), for example. The first clause says that the incidence of heavy feathers among feathers is higher than the incidence of heavy appleseeds among appleseeds. The second clause asserts that the incidence of heavy feathers among feathers is the same as the incidence of heavy appleseeds among appleseeds—to wit, nil. But that is a straightforward contradiction. And this prediction of a contradiction is unacceptable, since all of the examples in (24) can be true.

One might object that what counts as heavy depends on the context, and we can surely imagine contexts in which feathers count as heavy while appleseeds do not. So there is nothing wrong with saying that a randomly picked feather is more likely to be heavy than a randomly picked appleseed. This is true enough, but it is not enough to rescue the relative generic strategy. In order for this rebuttal to defuse my objection, it must accept that the contextually determined standard of heaviness changes between the interpretation of the two clauses. That's because the point of the objection is not that either of the clauses *feathers are heavier than appleseeds* or *neither feathers nor appleseeds are heavy* are unacceptable on their own. The point is that the relative generic proposal I'm considering is committed to saying that they are contradictory when they clearly are not. And the only way to avoid this prediction of contradictoriness is to accept that the interpretation of *heavy* in the two clauses changes. But in general, the context does not change in the course of interpreting the kind of sentence I am considering: *I am tall, though of course I am not tall* is simply contradictory.

That the present proposal founders on such cases should be unsurprising, since in general, a comparative can be true even though the corresponding positive is false.

- (25) a. Mary is taller than Sue, even though neither Mary nor Sue are tall.
- b. John is richer than Bill, even though neither John nor Bill are rich.

This discussion at least strongly suggests that we won't be able to extend the relative generic strategy to deal with generic comparisons.¹⁷

¹⁷It seems to me that this result also casts doubt on the viability of the relative generic strategy as it applies to its intended range of cases, such as (5). Given that the semantics of gradable predicates in their positive and comparative forms are very closely related, we should expect a proper treatment of their contribution to generic sentences to be uniform. Hence, any semantic theory for the positive case does not extend to the comparative case is therefore undermined.

3.3 Dependent Plurals

I now turn to a way of rejecting the second assumption, that both bare plurals are interpreted generically. The most reasonable way of implementing such a rejection is to take the object NP as a dependent plural of the kind illustrated in (26), due to Chomsky (1975).

(26) Unicycles have wheels.

Any one unicycle only has one wheel, so the plural morphology of the object NP *wheels* does not indicate that any one of the objects picked out by the subject has more than one wheel. Rather, it conveys roughly that between them, the unicycles have more than one wheel.¹⁸ So to a good approximation, a sentence with a dependent plural conveys that for each thing x denoted by the subject, there is at least one thing y denoted by the object such that x stands in the relation denoted by the predicate to y , and further, that the x s together stand in that relation to more than one y . In other words, we can essentially treat the bare plural in the object position as existentially quantified.

The point of the present proposal is perhaps best appreciated by contrasting it with the interpretations of (3) that resulted from interpreting both bare plurals in terms of a generic quantifier, repeated here as (27).

- (27) a. $[\text{GEN } x: \text{Girl}(x)] [\exists \theta] (\text{Does.Well}(x, \theta) \wedge$
 $[\text{GEN } y: \text{Boy}(y)] (\neg \text{Does.Well}(y, \theta)))$
 b. $[\text{GEN } x: \text{Girl}(x)] [\exists \theta] (\text{Does.Well}(x, \theta) \wedge$
 $\neg [\text{GEN } y: \text{Boy}(y)] (\text{Does.Well}(y, \theta)))$

The present proposal simply replaces the second generic quantifier with an existential one, as in (28a) and (28b).

- (28) a. $[\text{GEN } x: \text{Girl}(x)] [\exists \theta] (\text{Does.Well}(x, \theta) \wedge$
 $[\exists y: \text{Boy}(y)] (\neg \text{Does.Well}(y, \theta)))$
 b. $[\text{GEN } x: \text{Girl}(x)] [\exists \theta] (\text{Does.Well}(x, \theta) \wedge$
 $\neg [\exists y: \text{Boy}(y)] (\text{Does.Well}(y, \theta)))$

¹⁸For more detailed discussion, see e.g., Spector (2007) and Zweig (2008) and references therein.

These truth conditions are more extreme versions of the doubly-generic ones (27b) and (28) and hence face the same problems, only more so. In this case, (28a) is the weak member of the pair, saying that all normal (most) girls meet or exceed a threshold of scholastic achievement that at least one boy fails to meet or exceed—i.e., they all do better than the weakest boy. Again, this is far too weak. (28b), by contrast, is far too strong, since it says that all normal girls do better than any boy, i.e., better than even the strongest boy.

Consider now what happens when we weaken the dependent-plural analysis by combining it not with a strong generic quantifier but with Cohen’s weaker, relative generic operator. Since this weakens the predicted truth conditions, we need not consider (28a) with the relative generic operator. Focus instead on (28b). The relative interpretations of generics, recall, compares the likelihood that a randomly picked member of one set has a property of interest with the likelihood that a randomly picked member of some other set or sets (the alternatives) has that property. As before, I’ll assume that the initial set is the set of girls and that the alternative set is the set of boys. The property is somewhat complex: it is the property of meeting or exceeding a threshold that no boy meets or exceeds. Reflecting on this property shows us immediately that the incidence of this property among the set of boys is, by necessity, nil. Hence, the relative generic version of (28b) is true so long as the probability that a randomly picked girl does better than any boy is greater than 0, i.e., so long as the best student is a girl. In other words, on the relative generic version of (28b), the sentence is simply equivalent to the claim that the best student is a girl. These truth conditions are also clearly wrong, since they can be satisfied when a single girl is the best student, while all of the worst students are made up by the remainder of the girls.

Thus, even if we reject one or more of the three assumptions I’ve discussed so far, we still won’t have an empirically and theoretically viable treatment of generic comparisons.

4 Generics, Covers, and Comparatives

I now turn to my proposal. It rejects the assumption that the generic quantificational element is either a nominal determiner or an adverb of quantification. To introduce an alternative, I’ll discuss some phenomena in non-generic plu-

rals that are very similar to those observed for generic comparatives, along with a theoretical treatment in terms of covers.¹⁹

As I emphasized in the beginning, generic comparisons can be true even if it's not the case that GEN-many elements picked out by the subject term stand in the relation to GEN-many elements picked out by the object term. In my critical discussion of other proposals, I also drew attention to two other facts. First, it's possible for every sentence in the schema (17) to be false.

- (17) a. *As are F-er than Bs.*
 b. *Bs are F-er than As.*
 c. *As are (exactly) as F as Bs.*

The SANDWICH scenario I offered in the previous section is an instance of this pattern. All of the examples in (29) are false when evaluated as descriptions of this scenario.

- (29) a. *Girls do better than boys in grade school.*
 b. *Boys do better than girls in grade school.*
 c. *Girls do as well as boys in grade school.*

This also immediately shows that the whole distribution of girls and boys is relevant to determining the truth or falsity of these examples. It will not do to simply interpret *As are F-er than Bs* as requiring that the most *F* *As* are *F-er* than the most *F* *Bs*. In this case, it would amount to interpreting *boys do better than girls in grade school* as saying that the best boys do better than the best girls, and clearly, that condition is satisfied in the SANDWICH scenario, even though the original generic comparison is false.

I'll now argue that we see precisely corresponding phenomena when we consider the interpretation of comparatives in *non-generic* plurals. I'll use that observation to motivate a treatment of generic comparisons that is modeled very closely on the non-generic case.

4.1 Comparatives in Non-Generic Plurals

Just as in the case of generic comparisons, comparisons involving two plural definite descriptions can be true even when not all of the things denoted by

¹⁹For a thorough introduction to covers, see especially Gillon (1987, 1992); Schwarzschild (1994, 1996).

the one description stand in the relevant relation to all of the things denoted by the other. (30)-(32) illustrate the structure.

(30) The Hatfields are taller than the McCoys.

(31) The frigates are faster than the destroyers.

(32) The buses that ran today were emptier than the subways that ran today.

Each of these has an appropriately weak reading. (30) has a reading that requires only that the Hatfield men are taller than the McCoy men, the Hatfield women taller than the McCoy women, and so on. (31) is discussed by Schwarzschild, who observes that it's true in the following situation. The fleet has seen two model years, so that there are new and old frigates along with new and old destroyers. The fleet is deployed to two different areas, one in which a speedy response is important, one in which it is not, so that the newer frigates and destroyers are deployed to the former area and the older frigates and destroyers to the latter. (31) is true even if the advances in naval technology are such that the newer destroyers are faster than the older frigates, so long as within each area, the frigates are faster than the destroyers. Finally, (32) has a reading on which it's true if the buses were emptier than the subways running at the same time, even if the rush-hour buses weren't emptier than the late-night subways.

Note crucially that the interpretation of the data I am giving is consistent with an exhaustive interpretation of the plural NPs, i.e., that a predicate is truly predicated of a plural description only if it truly applies to all of the things denoted by it.²⁰ Brisson (2003) has suggested that in some cases, exhaustivity is suspended. She focuses on examples such as *the girls jumped in the lake*, which does not entail that every single one of the (contextually salient) girls jumped. However, trying to account for the data on comparatives by suggesting that exhaustivity is suspended is unpromising here. If the reason for the acceptability of (say) (31) was that the fast destroyers were irrelevant to the truth in the way that some girls are irrelevant to the truth of *the girls jumped in the lake*, then no matter what was the case with the irrelevant ships, the truth value of (31) should not change. But that is not what we find. If the fast destroyers were faster than the fast frigates, (31) would be clearly false. So

²⁰See, e.g., Fodor (1970); Sharvy (1980) for reasons for adopting this view.

just as in the case of the generic comparisons, members of both kinds across the whole distribution matter to the interpretation of the comparison.

The approach to these data that I will discuss makes use of the technical apparatus of *covers*, which we can introduce as a generalization of simple distributivity. When a plural NP is used distributively, it is used to summarize what a number of things did individually. This contrasts with collective uses, where the predicate that applies to the collection does not apply to each of its members. (33) illustrates both.

- (33) a. The children woke up at 8:00.
b. The children gathered in the yard.

(33a) entails that each of the children woke up at 8:00, while (33b) does not entail that each of the children gathered in the yard.

We could assume that distributivity is simply a brute feature of certain verbs, so that distributivity or collectivity is part of their lexical meaning. As several theorists have argued, a better account posits a distributive operator *D*.²¹ As a rough first approximation, we assume that a plural NP picks out a plurality and that the distributive operator applies to a VP. Assuming that *the children* picks out some contextually determined children, the LF of (33a) is given in (34a). (34b) gives the truth conditions we would like for (34a), and (34c) is the semantic value of the distributive operator. (I use capital variables to indicate that the variable ranges over pluralities, and the expression *Xx* to indicate that *x* is among the *Xs*.)

- (34) a. $[_S[_N \text{The children}][_V D \text{ woke up at eight}]]$.
b. $[\forall x: \text{Among.The.Children}(x)](\text{Woke.Up.At.Eight}(x))$
c. $[[D]] = \lambda f. \lambda X. [\forall x: Xx = 1](f(x) = 1)$

However, positing a simple distributive operator does not allow us to capture all of the semantic possibilities for interpreting sentences containing plural NPs, because there are sentences that aren't collective, but in which the predicate cannot be distributed all the way down to the individuals making up the plurality, either. Suppose we're buying apples. Each apple costs fifty cents,

²¹See, e.g., Beck and Sauerland (2000), Landmann (2000), Lasnik (1995), McKay (2006), Pietroski (2005), Schein (1993), Schwarzschild (1996), Winter (2000).

and we buy twelve. We can describe that situation accurately both by (35a) and (35b).²²

- (35) a. The apples cost fifty cents.
b. The apples cost six dollars.

We can predict both of these readings. In (35a), the VP contains the *D* operator, while it is absent in (35b). But if the apples come pre-wrapped in six-packs, we can also truly describe the situation with (36).

- (36) The apples cost three dollars.

We've already exhausted the possibilities with respect to the *D* operator in the LF of these sentences. When it's present, (35a) is true while (35b) and (36) are false. When it's absent, (35b) is true and the other two false. Either way, we cannot predict the true reading of (36). What we'd like is a formal way of capturing roughly the paraphrase (37) of (36).

- (37) The apples are such that six among them cost three dollars.

The predicate is not distributed to each of the objects picked out by the subject term, but rather to collections that in turn make up the collection picked out by the subject term, in this case, the pre-wrapped sixpacks.

That suggests that we should allow the distributive operator to distribute the predicate not just to the atoms, but to subclasses of the plurality denoted by the subject. Moreover, it seems as if the possibility of such intermediate readings depends heavily on the context. Without the information that the apples come in sixpacks, (36) is extremely hard to hear as true. Once we've made the sixpacks salient, the relevant reading becomes available. That suggests that the distributive operator should be sensitive to the context.

The key formal tool to accomplish all of these goals is that of a cover, which for our purposes we can simply treat as a partition.

[COVER] Where *A* is a set, *C* covers *A* iff

- (i) *C* is a set of subsets of *A*, and

²²Examples like these are due to the work of Roger Schwarzschild, as is the introduction of covers to account for them. See Schwarzschild (1994, 1996).

- (ii) $(\forall x \in A)(\exists! B)(B \subseteq A \wedge B \in C \wedge x \in B)$, and
- (iii) $\emptyset \notin C$.²³

That is, a set of sets C covers a set A just in case for every member of A , there is exactly one subset of A that contains that member, and that subset is in C . In the case of the apples, for example, A consists of the twelve apples, and the contextually salient cover C consists of two subsets, each of which contains the members of one sixpack.

We now need to incorporate covers into the semantics of the distributive operator. According to (34c), D takes a predicate as argument and returns another one. We now let D take two arguments, a predicate as before, as well as a contextually determined cover cov . Thus, the LF of (36) is given by (38a) and the truth conditions are (38b). On the assumption that the denotation of cov indeed is the set of two sets containing six apples each, these truth conditions are the ones we wanted to predict. Here and in what follows, I will treat it as a presupposition that the contextually supplied cover covers the collection picked out by the subject.²⁴

- (38) a. $[_S[_{NP}\text{The apples}][_ {VP}[D \text{ cov}] \text{ cost three dollars}]]$.
- b. Presupposition: $\llbracket \text{cov} \rrbracket$ covers $\llbracket \text{The apples} \rrbracket$
 Assertion: $[\forall y: y \in \llbracket \text{cov} \rrbracket] (\text{Cost.Three.Dollars}(y))$

More generally, the LF of a simple subject-predicate sentence with a non-generic plural is given in (39a), the corresponding truth conditions in (39b), and the semantic value of the revised D operator in (39c).²⁵

- (39) a. $[_S[_{NP}\text{The } A_{sPL}][_ {VP}[D \text{ cov}] F]]$
- b. Presupposition: $\llbracket \text{cov} \rrbracket$ covers $\llbracket A \rrbracket$
 Assertion: $[\forall y: y \in \llbracket \text{cov} \rrbracket] (\llbracket F \rrbracket (y) = 1)$

²³This is usually called a cover, rather than simply a partition because covers are supposed to allow that for some members of A , there is more than one subset in C , each of which contains that member. This is impossible for partitions. But that's a complication we can ignore, so that I'll remain with partitions.

²⁴Though the interpretation of the cover argument depends on the context, I'll suppress explicit reference to the context throughout for ease of exposition. Thus, in (38b), we are really evaluating $\llbracket \text{cov}_i \rrbracket^g$, i.e., an indexed cover variable that is assigned a value by the contextually salient assignment function (see Heim and Kratzer, 1998). But nothing I say will depend on being explicit about this.

²⁵I use the notation $\lambda \alpha : \beta. \gamma$ to express that the λ -abstract carries the presupposition that β is satisfied.

$$c. \llbracket D \rrbracket = \lambda c. \lambda f. \lambda X: c \text{ covers } X. \left[\left[\forall x: x \in c \right] (f(x) = 1) \right]$$

Distribution to the atoms is then simply the special case in which the cover consists of singleton sets, each of which contains just one of the members of the set picked out by the subject term.²⁶

We're very close to having an account of the example of the frigates and the destroyers. All we need to do is generalize the semantics presented thus far to sentences that contain a transitive verb and a plural noun phrase as a direct object. Following Schwarzschild, I assume that covers need not just be sets of sets. They can also be sets of pairs of sets. Formally, we introduce paired covers.

[**PAIRED COVER**] T is a *paired cover* of $\langle A, B \rangle$ iff

there is a cover of A , call it C_A , and there is a cover of B , call it C_B , such that

- (i) $T \subseteq C_A \times C_B$.
- (ii) $(\forall x \in C_A)(\exists y \in C_B)(\langle x, y \rangle \in T)$
- (iii) $(\forall y \in C_B)(\exists x \in C_A)(\langle x, y \rangle \in T)$

In order to accommodate paired covers in the semantics, I'll introduce a distributive operator D_T (the subscript T indicates that it is suitable for interpreting transitive verbs), which takes as one of its arguments such a paired cover.

We can then give the truth conditions of our initial target sentence (31), by assigning it the LF (31a) and the truth conditions (31b).

(31) The frigates are faster than the destroyers.

- a. $[_S [_{NP} \text{The frigates}] [_{VP} [[D_T \text{ cov}_T] \text{ are faster than}] \text{ the destroyers}]]$
 - b. Presupposition: $\llbracket \text{cov}_T \rrbracket$ covers $\langle \llbracket \text{the frigates} \rrbracket, \llbracket \text{the destroyers} \rrbracket \rangle$
- Assertion: $\left[\forall \langle p_1, p_2 \rangle: \langle p_1, p_2 \rangle \in \llbracket \text{cov}_T \rrbracket \right] \left(\text{Faster}(\langle p_1, p_2 \rangle) \right)$

²⁶A formal note: the proposal I'm discussing in the text has the result that a distributive reading of a sentence with a plural subject predicates the property denoted by the VP of a singleton set containing an atom making up the plurality. This analysis thus implicitly relies on the convention that we may conflate a singleton set with its member for the purposes of the semantics. This idea is due to Quine (2004).

The next issue concerns interpreting the nuclear scope in (31b), especially when the paired cover does not distribute all the way to the atoms of the pluralities denoted by the subject and object. Assume for concreteness that the old frigates are F_1, F_2, F_3 , the new frigates F_4 and F_5 , the old destroyers are D_1 and D_2 and the new destroyers D_3 and D_4 . Assume also that the contextually salient cover cov_T contains the pairs in (40).

$$(40) \llbracket \text{cov}_T \rrbracket = \{ \langle \{F_1, F_2, F_3\}, \{D_1, D_2\} \rangle, \langle \{F_4, F_5\}, \{D_3, D_4\} \rangle \}.$$

So in order to evaluate (31b), we have to evaluate (41).

$$(41) \text{Faster}(\langle \{F_1, F_2, F_3\}, \{D_1, D_2\} \rangle)$$

Here, two sets are compared. I propose that we introduce the following meaning postulate. When we compare two sets, the asserted relation holds between the two sets iff it holds between any pair of atoms from the two sets.

$$(42) \llbracket \text{comp} \rrbracket (\langle S_1, S_2 \rangle) = 1 \text{ iff} \\ [\forall x: x \in S_1] [\forall y: y \in S_2] (\llbracket \text{comp} \rrbracket (\langle x, y \rangle) = 1)$$

This postulate requires that each of F_1, F_2 , and F_3 is faster than any of D_1 or D_2 . The motivation for adopting this assumption comes from reflecting on the following example from Schwarzschild (1996, 86). Consider table 1, which can be truly described by (43).

Fiction	Non-fiction
Alice in Wonderland	Aspects Language (Bloomfield)
Fantastic Voyage	Gray's Anatomy
David Copperfield	Das Kapital
Hard Times	The Wealth of Nations
Oedipus Rex	Freud's Introduction to Psychology
Agamemnon	
Richard III	Machiavelli's The Prince

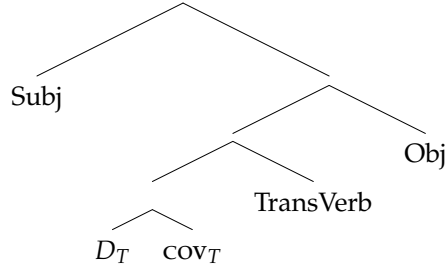
Table 1: Books

(43) The fiction books complement the non-fiction books.

The true reading comes about when we assume that a pair-cover is contextually salient which pairs books from the same line. Crucially, (43) becomes false if we add a non-corresponding book to any of the lines. That's why I suggest we impose the demanding condition that the relation holds between every pair that can be formed from the two sets.

We can now state the general view, where the schematic LF (44a) is assigned the truth conditions (44b), and the semantic value of the transitive distributive quantifier is given in (44c).

(44) a.



b. Presupposition: $\llbracket \text{COV}_T \rrbracket$ covers $\langle \llbracket \text{Subj} \rrbracket, \llbracket \text{Obj} \rrbracket \rangle$

Assertion: $\left[\forall \langle p_1, p_2 \rangle : \langle p_1, p_2 \rangle \in \llbracket \text{COV}_T \rrbracket \right]$
 $\left(\llbracket \text{TransVerb} \rrbracket (\langle p_1, p_2 \rangle) = 1 \right)$

c. $\llbracket D_T \rrbracket = \lambda c. \lambda f. \lambda X. \lambda Y : c \text{ covers } \langle X, Y \rangle.$

$\left[\left[\forall \langle p_1, p_2 \rangle : \langle p_1, p_2 \rangle \in c \right] (f(\langle p_1, p_2 \rangle)) \right]$

In this presentation I am skirting over several large issues, in particular, over how distributivity interacts with comparative morphology to compositionally determine the interpretation I am giving here.²⁷ What is crucial for my purposes is that the distributive operator applies to the predicate before it applies to the subject and object.²⁸ Since the relevant readings are clearly attested, and this assumption seems to be quite clearly required in order to derive them, I'll remain at this level of generality.

But even at this level of description, we can already see that the semantics for non-generic comparisons predicts the following fact about logical relations. Consider the schemata in (45).

²⁷For relevant discussion, see Fitzgibbons et al. (forthcoming) and references therein.

²⁸I assume that this occurs via movement of the subject and object, but I'll suppress the interpretation of the movement. For discussion, see Beck and Sauerland (2000); Sauerland (1998); Sternefeld (1998).

- (45) a. The *As* are *F*-er than the *Bs*.
 b. The *Bs* are *F*-er than the *As*.
 c. The *As* are (exactly) as *F* as the *Bs*.

The semantics I've just presented immediately predicts that all of these can be false. The basic point is just that the three schemata in (45) do not jointly exhaust the space of possibilities. Suppose that we have a paired cover over the *As* and the *Bs*. It may be true that for some pairs in that cover, the *As* in that pair are *F*-er than the *Bs* while for other pairs, the *Bs* are *F*-er than the *As*. In that situation, all of the schemata in (45) are false.

And crucially, the semantics makes that prediction precisely because the *As* and the *Bs* throughout the distribution of *F*-ness are relevant to determining the truth of (45a)-(45c), not just the *As* and the *Bs* that are *F*-est, or the average *A* and the average *B*.

4.2 Covers in Generic Comparisons

The fact that we see such a close parallel in the logical relations among non-generic comparatives involving plurals and generic comparisons suggests that we can make progress on the latter by adapting the semantics of the former. That's what I'll do right now.

In the case of the non-generic plurals, the basic quantificational force derives from a distributive operator, so the most direct way of transferring this mechanism is to treat the generic operator not as a nominal determiner nor as an adverb of quantification, but as a distributive operator. Moreover, the analogy with the universal distributive operator in the non-generic case suggests that we treat the generic operator as a universal quantifier of some stripe.

Clearly, characterizing sentences are not well-paraphrased as universally quantified claims, since they are compatible with what would be counterexamples to the corresponding universal claim. Hence, we should interpret the generic operator as a restricted universal quantifier. For the purposes of this paper, I'll say that the restriction is to the normal members of a kind. It's obviously extremely difficult to say what makes a member of a kind normal in any kind of general way, but when we are confronted with a particular generic, we have a clear enough sense of which members of the kind are intuitively relevant to its truth and which ones are not. For example, when we're evaluating

(3), we're not interested in academic outliers. We have a notion of a normal student, girl or boy, and the generic is evaluated with respect to those girls and boys, not the ones that are exceptionally good or exceptionally bad.

Quantificational accounts generally need to make the restriction of the quantifier sensitive to the predicate of the generic. The need for this is most easily brought out by considering the pair of sentences (46).

- (46) a. Chickens lay eggs.
- b. Chickens are hens.

If the restriction of the quantifier was the same in both (46a) and (46b), then (46a) would entail (46b). After all, since a chicken lays eggs only if it is a hen, (46a) requires that all of the chickens in the scope of the quantifier are hens, and if the scope of the quantifier in (46b) was the same, it would follow that chickens are hens. Informally, we might say that what is at issue is not being a normal chicken, but being a chicken that is normal in a respect, and that the respect is determined by the predicate of the characterizing sentence. Thus, (46a) might be about all of the chickens that are normal with respect to how chickens extrude offspring (and no male qualifies for being normal in this respect), while (46b) might be about all of the chickens that are normal with respect to having a gender (which includes both female *and* male chickens).²⁹

It is also true that generics can be non-vacuously true, even when there are no normal members of the kind at issue in the world of evaluation at the time of evaluation. To take a simple example, *lions have four legs* can be non-vacuously true even when there aren't any normal lions, perhaps because all of the lions have lost a leg in accidents or fights. That means that when we evaluate a generic, we always need to ensure that we evaluate the restricted universal quantifier with respect to a suitably representative domain. One way to accomplish this is to introduce a counterfactual element into the semantics. Simple generic sentences of the form (47a) are thus interpreted as (47b).

- (47) a. *As* are *B*.
- b. If there was a relevantly normal *A*, then all relevantly normal *As* would be *B*.

²⁹Cohen implements the same strategy by using a set of alternatives to the predicate in order to restrict the domain of his generic operator.

Here, *relevantly normal* just abbreviates *normal in the respect determined by the predicate*. If there are relevantly normal As in the world of evaluation at the time of evaluation, then (47b) just collapses into a simple quantificational claim. Moreover, we see a useful interplay between the fact that we need to consider not normal As, but relevantly normal As, since that helps make the counterfactual strategy work. Suppose, for example, that we were interested in evaluating *lions have manes* and analyzed it as *if there were any normal lions, then all normal lions would have manes*. We would not have any grounds for denying that female lions are normal lions. Hence, we would predict that *lions have manes* is false. But by restricting the generic quantifier to those lions that are normal with respect to their ornamentation, we may exclude the female lions.

The rest of this section is devoted to combining this basic semantic theory for non-comparative generics with the semantics for non-generic comparatives. Return to (3), *girls do better than boys in grade school*, for illustration. Suppose that we've decided on the relevantly normal boys and girls, excluding the outliers, and suppose that the SHIFT-scenario obtains.

Here is a way of assigning truth conditions to (3) that makes the correct prediction. We match up subsets of the girls with subsets of the boys, where each of the subsets comprises a certain part of the distribution, for example pairing deciles of girls with corresponding deciles of boys. The sentence is true iff within each such pairing, each of the girls does better than any of the boys. Graphically, we can represent this strategy as in figure 2, where areas that are shaded the same way are compared to each other.

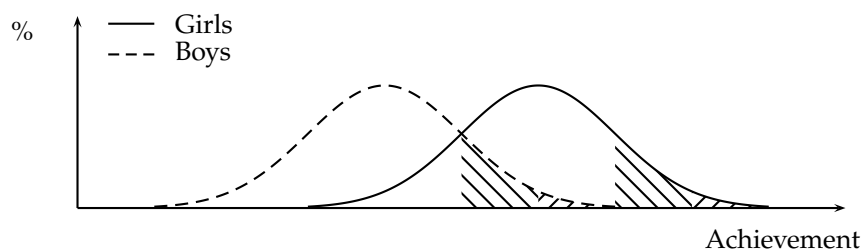


Figure 2: Correspondences for SHIFT

That is the informal idea I now give a formally more rigorous implementation of. I'll go stepwise, beginning with simple sentences and adding com-

plexity as I go along. If *gen* occupies the position of a distributive operator, and ignoring covers for now, the LF of the simple characterizing sentence (1), repeated here as (48), is (48a).

(48) Ravens are black.

$$\text{a. } [_S [_{NP} \text{ ravens }] [_{VP} [\text{gen } \text{black}]]]$$

Thus, I want to interpret the LF (48a) as having the truth conditions (48b), which we can achieve compositionally if *gen* has the lexical semantics in (48c).

$$(48) \quad \text{b. } [\forall x : \text{Relevantly.Normal}(x) \wedge \text{Raven}(x)] (\text{Black}(x)).$$

$$\text{c. } \llbracket \text{gen} \rrbracket = \lambda f. \lambda g. \left[[\forall x : \text{Relevantly.Normal}(x) \wedge g(x)] (f(x) = 1) \right]$$

We can immediately introduce covers. Just as the distributive operator takes an extra cover argument, the generic quantifier *gen* does, too. The only important change we have to make is to alter the interpretation of the subject term to be more in line with the denotation of non-generic plurals. Rather than have the subject term pick out a property, it should pick out the plurality of members of the kind. The relevant LFs, truth conditions, and lexical semantics are given in (49).

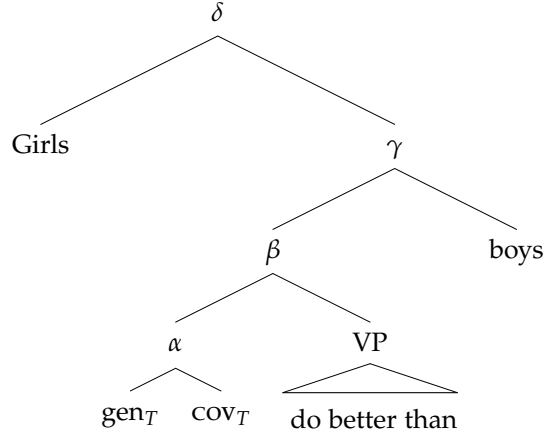
$$(49) \quad \text{a. } [_S [_{NP} \text{ subj }] [_{VP} [\text{gen } \text{cov }] \text{pred }]]$$

$$\text{b. } \llbracket \text{gen} \rrbracket = \lambda c. \lambda f. \lambda X : c \text{ covers } X. \left[[\forall x : \text{Relevantly.Normal}(x) \wedge x \in c] (f(x) = 1) \right]$$

As before, quantification over the atoms in the set picked out by the subject term falls out as a special case.

To account for the interpretation of generic comparisons, I make use of paired covers. Thus, I want to offer roughly the LF (50a) as the interpretation of (3), and the semantics for the nodes in the tree in (??).

(50) a.



- b. $\llbracket Girls \rrbracket = \{x: x \text{ is a girl}\}$
 $\llbracket Boys \rrbracket = \{y: y \text{ is a boy}\}$
 $\llbracket do \text{ better than} \rrbracket = \lambda Y. \lambda X. [X \text{ does better in grade school than } Y]$
 $\llbracket gen_T \rrbracket = \lambda c. \lambda f \lambda Y. \lambda X: c \text{ covers } \langle X, Y \rangle. [\forall \langle p_1, p_2 \rangle: \langle p_1, p_2 \rangle \in c \wedge$
 $\text{Relevantly.Normal}(p_1) \wedge \text{Relevantly.Normal}(p_2)] (f(p_2)(p_1) = 1)$
 $\llbracket$

The derivation of (3), given this information, is the following.

[DERIVATION OF [3]]

$$\begin{aligned}
 \llbracket \alpha \rrbracket &= \lambda f \lambda Y. \lambda X: \llbracket COV_T \rrbracket \text{ covers } \langle X, Y \rangle. \\
 &\quad [\forall \langle p_1, p_2 \rangle: \langle p_1, p_2 \rangle \in \llbracket COV_T \rrbracket \wedge \\
 &\quad \text{Relevantly.Normal}(p_1) \wedge \text{Relevantly.Normal}(p_2)] (f(p_2)(p_1) = 1) \\
 \llbracket \beta \rrbracket &= \lambda Y. \lambda X: \llbracket COV_T \rrbracket \text{ covers } \langle X, Y \rangle. [\forall \langle p_1, p_2 \rangle: \langle p_1, p_2 \rangle \in \llbracket COV_T \rrbracket \wedge \\
 &\quad \text{Relevantly.Normal}(p_1) \wedge \text{Relevantly.Normal}(p_2)] \\
 &\quad (p_1 \text{ does better in grade school than } p_2) \\
 \llbracket \gamma \rrbracket &= \lambda X: \llbracket COV_T \rrbracket \text{ covers } \langle X, \{y: y \text{ is a boy}\} \rangle. \\
 &\quad [\forall \langle p_1, p_2 \rangle: \langle p_1, p_2 \rangle \in \llbracket COV_T \rrbracket \wedge \\
 &\quad \text{Relevantly.Normal}(p_1) \wedge \text{Relevantly.Normal}(p_2)] \\
 &\quad (p_1 \text{ does better in grade school than } p_2) \\
 \llbracket \delta \rrbracket &= \text{Presupposition: } \llbracket COV_T \rrbracket \text{ covers } \langle \{x: x \text{ is a girl}\}, \{y: y \text{ is a boy}\} \rangle \\
 &\quad \text{Assertion: } [\forall \langle p_1, p_2 \rangle: \langle p_1, p_2 \rangle \in \llbracket COV_T \rrbracket \wedge \\
 &\quad \text{Relevantly.Normal}(p_1) \wedge \text{Relevantly.Normal}(p_2)] \\
 &\quad (p_1 \text{ does better in grade school than } p_2)
 \end{aligned}$$

That is to say, for every pairing in the contextually salient cover, every girl in the pairing does better in grade school than every boy in that pairing. Given that the contextually salient cover pairs up girls and boys in the same decile in the distribution, we have the following interpretation of (3): it is true just in case every girl in the top decile of girls does better than every boy in the top decile of boys, and so on down the distribution of scholastic achievement.

This proposal requires some further comments before I discuss its benefits. First, one might worry that these truth conditions are too demanding. For concreteness, consider a situation in which all but one girl is better than all of the boys, but in which the weakest student is a girl. One might think that my semantics predicts that (3) is false—after all, this is just an extreme SANDWICH-style scenario—but that intuitively, (3) is true.

This objection is too quick, and to see why, it will help to spell out the situation further. Different ways of spelling it out will yield different intuitive truth-value judgments, and my proposal is in accord with these judgments. The crucial question is whether the very weak girl is normal for the purposes of evaluating (3). Perhaps we are thinking about (3) in the context of designing social policy, or wondering why women aren't more heavily represented in many academic disciplines. So now suppose that the weakest girl is doing poorly in school because of a learning deficit. In that case, she is not normal for the purposes of evaluating (3), so that she is not in the scope of the generic operator and my semantics makes the intuitively correct prediction that (3) is true. Suppose, however, that this girl is in no relevant way different from the girls at the top, aside from her level of achievement—there is nothing we could point to that certifies her as an outlier. In that case, my semantics predicts that (3) is indeed false.

A potentially confounding factor may be the fact that whenever we see a single girl or a single boy in a chunk of the distribution, we naturally assume that this child is an outlier. After all, there are so many kids that generally shared factors should produce the same outcome in multiple cases. That's why in the extreme situation I just described, the easiest judgment is that (3) is true. And the prediction of my semantics conforms to that judgment.³⁰

³⁰An anonymous referee has emphasized a conflicting intuition regarding this example. Even if the weakest student is a girl and she is completely normal, it can be true that girls do better than boys in grade school. I can see at least two possible responses to this observation. The first is that comparative generics are simply ambiguous between the reading I have been focusing on and a

One may, of course, also worry about my appeal to an unanalyzed notion of normality in the semantics. And it is certainly true that it can be hard to determine what exactly counts as normal for the purposes of evaluating (3): what factors determine whether a distribution is normal for our purposes? But I take it to be an advantage of my account that it leaves these questions open. After all, it happens quite frequently, especially when we're considering topics like gender differences, that speakers agree on the statistical facts but disagree on the generic facts. I diagnose this as a disagreement over what counts as relevantly normal.

Finally, let me briefly address what happens when we do not have a representative domain. Because of some random fluke, a part of the distribution that is normally occupied by some girls is completely vacant. Perhaps it just so happened that all of the high-performing girls moved away. It is in order to deal with this problem that I suggest that the cover itself be intensional. That allows me to say the following: it has to be the case for each pair $\langle A_i, B_i \rangle$ in the paired cover over the normal *As* and the normal *Bs* that, if there were *As* and *Bs* in the cells A_i and B_i , then each of the *As* in A_i would be *F*-er than any of the *B* in B_i . For the special case of (3): for each pair $\langle \text{Girls}_i, \text{Boys}_i \rangle$ in the contextually selected paired cover over the normal girls and normal boys, if there were any girls in Girls_i and any boys in Boys_i , then each of the girls in Girls_i would do better in grade school than any of the boys in Boys_i . In other words, the counterfactual strategy for dealing with non-representative domains I introduced above for non-comparative generics can be directly extended here.

In §3.1, I observed that the logical relations among generic comparisons do not mirror those of comparatives involving individual objects. Schematically, all of the sentences in (17) can be false.

- (17) a. *As* are *F*-er than *Bs*.
 b. *Bs* are *F*-er than *As*.
 c. *As* are (exactly) as *F* as *Bs*.

reading in terms of averages. The second is that, once we have relatively large groups, we allow for a bit of imprecision in our speaking. Regarding the second point, notice that how good (3) is in a situation in which the weakest student is a girl seems to vary with how large the population of students is: the more students, the easier it is to disregard the weakest girl as somehow not standing in the way of (3)'s truth. I leave as an open problem how this kind of situation can be handled.

My semantics predicts that this is possible without further assumptions. I'll give the argument for the particular example of (3). That is, I'll show that the three claims in (29), repeated from earlier, can all be false in the SANDWICH scenario.

- (29) a. Girls do better than boys in grade school. [= (3)]
- b. Boys do better than girls in grade school.
- c. Girls do as well as boys in grade school.

Given the semantics I worked out for (29a), it is true iff within each decile, every girl does better than every boy. Since there are some deciles in which each of the boys do better than the girls—the ones at the top—it is false that girls do better than boys. By parallel reasoning, (29b) is true iff within each decile, every boy does better than every girl. Since there are some deciles in which each of the girls does better than any of the boys—the ones at the bottom—it is false that boys do better than girls. Finally, (29c) is true iff within each decile, the girls do as well as the boys. It doesn't matter how exactly we make precise the notion of one group's doing as well as another. For concreteness, assume that two groups do equally well when the top girl does as well as the top boy in that group, and the bottom girl does as well as the bottom boy in that group. However we choose to settle the matter, it will always turn out that some of the deciles fail that condition, for example, the top deciles in which the boys do far better than the girls.

4.3 Dimensions of Contextual Variation

My semantics combines two distinct semantic features to yield interpretations for generic comparisons, a univocal generic operator that I'm glossing as a restricted universal quantifier and a contextually selected paired cover. Since both of these features are present in sentences besides generic comparisons, we should be able to find commonalities between generic comparisons and those sentences in which one, but not the other, of the features is present. I want to argue that this is exactly what we do find, focusing on kinds of contextual variability, beginning with variability in the contextually determined cover and then turning to contextual variability in what counts as normal.

Consider the two sentences in (51).

- (51) a. Men are taller than women.
 b. Fathers are taller than mothers.

On its most salient reading, (51a) is true. We simply compare the distribution of men and women with respect to height and pair up occupants of corresponding parts. By contrast, (51b) has two natural readings, one on which it is true, another on which it is false. The true reading is the one that is parallel to the reading of (51a) I just described: we compare the tallest men who happen to be fathers with the tallest women who happen to be mothers, the less tall such men with the less tall such women, and so on down the distribution. However, (51b) also has a very natural reading on which we don't pair the tallest men who happen to be fathers with the tallest women who happen to be mothers, but instead pair each father with the mother of their common child or children. In that case, (51b) is false because there are plenty of pairings consisting of fathers and mothers, both of whom fall in the respective normal height ranges where the father is shorter than the corresponding mother. We can make this latter reading salient by mentioning this kind of pairing, as in (52).

- (52) Couples with kids can be very varied—fathers (certainly) aren't taller than mothers.

The very same phenomenon is present in non-generic plurals, as (53) shows.

- (53) Even though the couples in our study were not married, the men did display aggressive behavior towards the women. (Schwarzschild, 1996, 87, #209)

The concessive in the first part of the sentence makes salient a pairing of men and women according to the couples that were formed, and the main clause is interpreted with respect to this pairing. If we remove the concessive, this reading becomes harder to hear.

- (54) In our study, the men displayed aggressive behavior towards the women.

It is clear that there is some matching of men with women, though for all that the context provides, it could be that every man in the study displayed the behavior towards every woman in the study.

Turn now to contextual variation in what counts as normal. Consider (55).

(55) Dogs are heavier than cats.

It is clear that, at the top of the distribution, dogs are far heavier than cats. Indeed, this is true for most parts of the size-distribution. The interesting aspect of the case is located at the bottom of the distribution, specifically when we consider so-called tea-cup dogs, breeds of dogs that are lighter than even the lightest cats. These tea-cup dogs exist only because of very targeted selection by breeders, and individuals of these species usually suffer from health problems and various genetic disorders. Given this information, the sentence (55) has at least a true and a false reading, depending on whether we take tea-cup breeds into account. Put in terms of the proposal I am putting forth, the difference between the context in which (55) expresses a truth and one in which it expresses a falsehood is simply a matter of whether such intervention by breeders disqualifies a dog from being normal.

The point I want to draw attention to is that the very same kind of contextual variability can be seen in non-comparative generics, as the texts in (56a) and (56b) show. Some background: at birth, dobermans have floppy ears and long tails. In many countries, including the US, breeders then surgically insert posts into the ears that remain there for about six weeks until the cartilage in the ears hardens, at which point the posts are removed. They also dock the dogs' tails.

- (56) a. Different breeds of dogs focus on different senses. Some dogs have very acute hearing, while others have a specialized sense of smell. The latter have floppy ears that agitate the air when they're following a trail. This is true of dobermans: *they don't have pointy ears, they have floppy ears.*
- b. Welcome to the Westminster Kennel Club show. We have a wide range of dogs, some homely, some truly regal. The dobermans are beautiful, dynamic creatures. *Dobermans have pointy ears. They don't have floppy ears.*

In each context set up by the discourses preceding the italicized sentences, we can deny a sentence that is asserted in the other. That shows that the difference between (56a) and (56b) really is a difference in the proposition expressed, and not just a pragmatic phenomenon that doesn't influence semantics. We

can thus observe kinds of contextual variation in sentences other than generic comparisons.

The discussion of these examples illustrates a general strategy for testing the proposal I've made. If we see contextual variability in the interpretation of generic comparisons, we should be able to trace it to an aspect of their interpretation that they share with either non-generic plural comparisons or non-comparative generics. These examples suggest that this strategy can, in fact, be carried out.

5 Conclusion

I've argued that generic comparisons exhibit a range of interpretations that are problematic for extant semantic treatments. I've also argued that these problems can all be solved in a theoretically motivated way if we treat generics not as quantificational constructions or instances of direct reference to kinds, but as plural constructions making use of a special generic distributive operator. Clearly, such a reconfiguration of the LF of generics will have many other implications, a fuller exploration of which I leave to further work.

References

- Asher, N. and M. Morreau. "What Some Generic Sentences Mean". In G. N. Carlson and F. J. Pelletier, eds., *The Generic Book*, 300–339 (Chicago: University of Chicago Press, 1995).
- Beck, S. and U. Sauerland. "Cumulation is Needed: A Reply to Winter (2000)". *Natural Language Semantics*, 8(4), (2000), 349–371.
- Brisson, C. "Plurals, All, and the Nonuniformity of Collective Predication". *Linguistics and Philosophy*, 26(2), (2003), 129–184.
- Carlson, G. and F. J. Pelletier. "The Average American has 2.3 Children". *Journal of Semantics*, 19(1), (2002), 73–104.
- Carlson, G. N. *Reference to Kinds in English*. Ph.D. thesis, University of Massachusetts, Amherst (1977).
- Chomsky, N. "Questions of Form and Interpretation". In R. Austerlitz, ed., *The Scope of American Linguistics*, 159–196 (Lisse: Peter de Ridder Press, 1975).
- Cohen, A. "Generics, Frequency Adverbs, and Probability". *Linguistics and Philosophy*, 22, (1999a), 221–253.
- . *Think Generic!* (Stanford, CA: CSLI Publications, 1999b).
- . "Relative Readings of *Many*, *Often*, and Generics". *Natural Language Semantics*, 9(1), (2001), 41–67.
- . "Generics and Mental Representation". *Linguistics and Philosophy*, 27(5), (2004), 529–556.
- Diesing, M. *Indefinites* (Cambridge: MIT Press, 1992).
- Fitzgibbons, N., Y. Sharvit, and J. Gajewski. "Plural Superlatives and Distributivity". In *Proceedings of SALT XVIII* (forthcoming).
- Fodor, J. D. *The Linguistic Description of Opaque Contexts*. Ph.D. thesis, MIT (1970).
- Gillon, B. "The Readings of Plural Noun Phrases". *Linguistics and Philosophy*, 10(2), (1987), 199–220.
- . "Towards a Common Semantics For English Count and Mass Nouns". *Linguistics and Philosophy*, 15(6), (1992), 597–640.
- Heim, I. and A. Kratzer. *Semantics in Generative Grammar* (Malden, MA: Blackwell, 1998).
- Kennedy, C. *Projecting the Adjective* (New York: Garland, 1999).

- . “Vagueness and Grammar: The Semantics of Relative and Absolute Gradable Adjectives”. *Linguistics and Philosophy*, 30(1), (2007), 1–45.
- Kennedy, C. and J. Stanley. “On ‘Average’”. *Mind*.
- Krifka, M., F. J. Pelletier, G. N. Carlson, A. ter Meulen, G. Chierchia, and G. Link. “Genericity: An Introduction”. In G. N. Carlson and F. J. Pelletier, eds., *The Generic Book*, 1–124 (Chicago: University of Chicago Press, 1995).
- Landmann, F. *Events and Plurality* (Dordrecht: Kluwer Academic, 2000).
- Laserson, P. *Plurality, Conjunction, and Events* (Dordrecht: Kluwer, 1995).
- Lewis, D. K. “Adverbs of Quantification”. In E. L. Keenan, ed., *Formal Semantics of Natural Language*, 3–15 (Cambridge, UK: Cambridge UP, 1973).
- McKay, T. J. *Plural Predication* (Oxford: Oxford UP, 2006).
- Pelletier, J. and N. Asher. “Generics and Defaults”. In J. van Benthem and A. ter Meulen, eds., *Handbook of Logic and Language*, 1125–1177 (Amsterdam: Elsevier and MIT Press, 1997).
- Pietroski, P. M. *Events and Semantic Architecture* (Oxford: Oxford UP, 2005).
- Quine, W. *Set Theory and Its Logic (revised edition)* (Cambridge, MA: Belknap Press, 2004).
- Sauerland, U. “Plurals, Derived Predicates, and Reciprocals”. In U. Sauerland and O. Percus, eds., *The Interpretive Tract*, 177–204 (Cambridge, MA: MITWPL, 1998).
- Schein, B. *Plurals and Events* (Cambridge, MA: MIT Press, 1993).
- Schubert, L. K. and F. J. Pelletier. “Generically Speaking, or, Using Discourse Representation Theory to Interpret Generics”. In G. Chierchia, B. H. Partee, and R. Turner, eds., *Properties, Types, and Meaning, Vol. II*, 193–268 (Dordrecht: Kluwer Academic Publishers, 1989).
- Schwarzschild, R. “Plurals, Presuppositions, and the Sources of Distributivity”. *Natural Language Semantics*, 2, (1994), 201–248.
- . *Pluralities* (Dordrecht: Kluwer, 1996).
- . “The Semantics of Comparatives and Other Degree Constructions”. *Language and Linguistics Compass*, 2(2), (2008), 308–331.
- Sharvy, R. “A More General Theory of Definite Descriptions”. *The Philosophical Review*, 89(4), (1980), 607–624.
- Spector, B. “Aspects of the Pragmatics of Plural Morphology: On Higher-Order Implicatures”. In U. Sauerland and P. Stateva, eds., *Presupposition and Implicature in Compositional Semantics*, 243–281 (Houndsmills: Palgrave Macmillan, 2007).

- Sternefeld, W. "Reciprocity and Cumulative Predication". *Natural Language Semantics*, 6(3), (1998), 303–337.
- Wilkinson, K. *Studies in the Semantics of Generic Noun Phrases*. Ph.D. thesis, University of Massachusetts, Amherst (1991).
- Williamson, T. *Vagueness* (London: Routledge, 1994).
- Winter, Y. "Distributivity and Dependency". *Natural Language Semantics*, 8(4), (2000), 27–69.
- Zweig, E. *Dependent Plurals and Plural Meaning*. Ph.D. thesis, New York University (2008).